

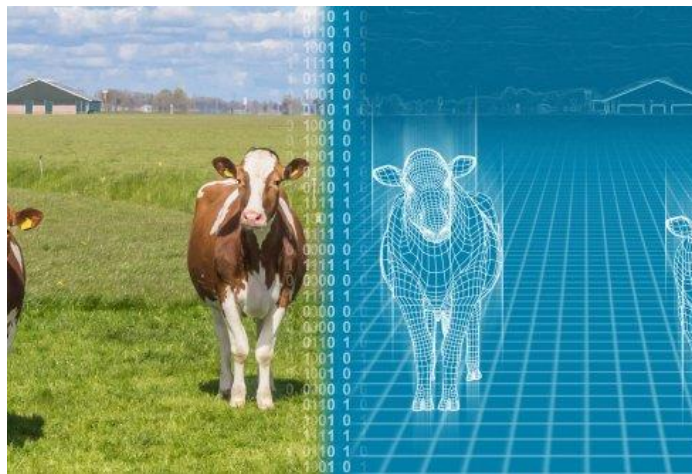
RESOLUTION AND DIGITAL TWINS

Working with complex real-world phenomena in a digital environment.

Deliverable of the project “Data Quality dimensions and communication”

Exploratory project in the Digital Twins Platform Methodology

Digital Twins investment theme



January, 2021

Jandirk Bulens, Martin Knotters, Wies Vullings, Dennis Walvoort, (Wageningen Environmental Research),
Ron Wehrens, (Wageningen Plant Research, Biometris)

Open Access

CC-BY-SA

VISIT...

LANZAROTE
Caliente.COM

Introduction

Digital Twins are a hot research topic. They originate from manufactured artifacts¹, but have recently made inroads in other communities such as the geographic information community, propelled by technical developments closing the gap between the physical environment and their digital representations². Also, within the WUR domain, digital twins are being developed representing all kinds of phenomena in our research domain: the three WDCC ‘flagship’ projects model a tomato plant, farm and a human digestive tract. At WENR work has been done on modelling a watershed. The big difference, though, between manufactured artifacts and phenomena such as a farm or a watershed is that the artifacts are complicated items and phenomena are complex. Complicated problems and systems originate from causes that each can be individually distinguished and for each input to the system there is a proportionate output. On the other hand, complex problems and systems result from networks of multiple interacting causes that cannot be individually distinguished. They must be addressed as entire systems, *i.e.*, they cannot be addressed in a piecemeal way³.

There is no single definition of what a Digital Twin is. However, it is possible to define a series of generalisations, or rules, from which we can establish what is a Digital Twin and what is not⁴. Succinctly put, (1) a Digital Twin must have a physical part and a virtual replica of it (2) between which a bidirectional or two-way flow of information should occur that (3) enables or initiates some feedback or decision-making processes. Beyond these generalisations, definitions vary quite a lot. Therefore, digital twins can have very different characteristics depending on their purpose, design, and goals⁵⁶.

A general definition and explanation from The Open Data Institute⁷ that seems to be in agreement with the series of generalisations is: ‘A Digital Twin refers to the virtual replication of a physical ‘thing’ or ‘phenomenon’. This replication is updated in real time (or as regularly as possible) to match its real-world counterpart as closely as possible. A Digital Twin can then be used for a variety of purposes, from testing new strategies to analysing previous results. In general sensors and other components are used to gather accurate data about the asset in question. These data are then processed and used to create the replica. Needless to say, more advanced sensors, whether in the form of more accurate equipment or a denser, more compact array, will help towards creating a more detailed digital model’.

A Digital Twin is seen as the ‘next thing after the Internet of Things’⁸: the IoT enables this through sensors and input data providing favourable conditions for digital twins. With more interconnected devices and applications, it’s easier for organisations to interlink data and get a better view of the wider picture. A Digital Twin, then, is simply an advanced way of collecting and combining these data for specific purposes.

Since a key feature of a Digital Twin is the ability to provide answers to questions that are not yet known, a careful set up is required. One of the most important factors is the level of detail or resolution that is required for the abstraction of the real-world case. We define resolution as the distribution of shortest distances between objects. In the spatial domain, examples related to resolution include a measure of the smallest object that can be resolved by the sensor, the ground area imaged by the instantaneous field of view (IFOV) of the sensor, or the linear dimension on the ground represented by each pixel⁹. For time, resolution is related to the distances in time between the measuring points (regular or not), and in the semantic domain by the level of detail in the categorization used.

In general, data are collected for a specific resolution. Since Digital Twins use data in an integrated manner, there are a number of pertinent questions. What is the resolution needed for the Digital Twin? What is the resolution of the available data? Are all data aligned? Is resolution in balance with data quality? What are implications for the results and the presentation and visualization of these resolution

issues? Processes modelled by the Digital Twin often can be described in various resolutions, e.g., a Digital Twin of a watershed will not only represent the watershed as a whole, but will also focus on the streams, ditches and canals within it. It will be interesting to look at the discharge of water after a large shower, but also at the effects of drought over the years.

A Digital Twin as a replica of a real-world thing or phenomenon must fit in a digital environment in its proper context, also in a fully conceptual way. Therefore, conceptual modelling is the process to create an abstract description of a part of the real world and / or a series of related concepts. Conceptual models may only exist in the minds of people who communicate them verbally and often inaccurately. They may also be written down and stored for wider dissemination¹⁰. In the ISO standards for geographic information conceptualisation starts with identifying the Universe of Discourse: a selected piece of the real world that a human being wishes to describe in a model. The selected piece is identified by the 'smallest' object of information relevant within that Universe of Discourse. The universe of discourse may include not only objects such as watercourses, lakes, islands, property boundaries, property owners and exploitation areas, but also their attributes, their functions and the relationships that exist among such features. The figure below describes the relationship between modelling the real world and the resulting conceptual schema.¹¹

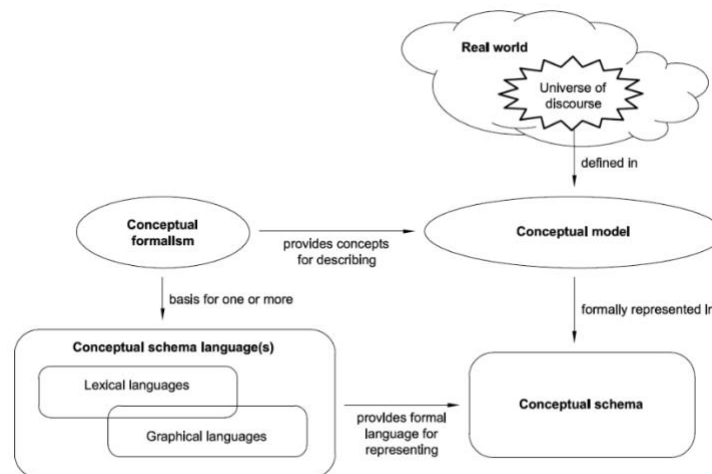


Figure 1 conceptual modelling the real world (ISO19101)

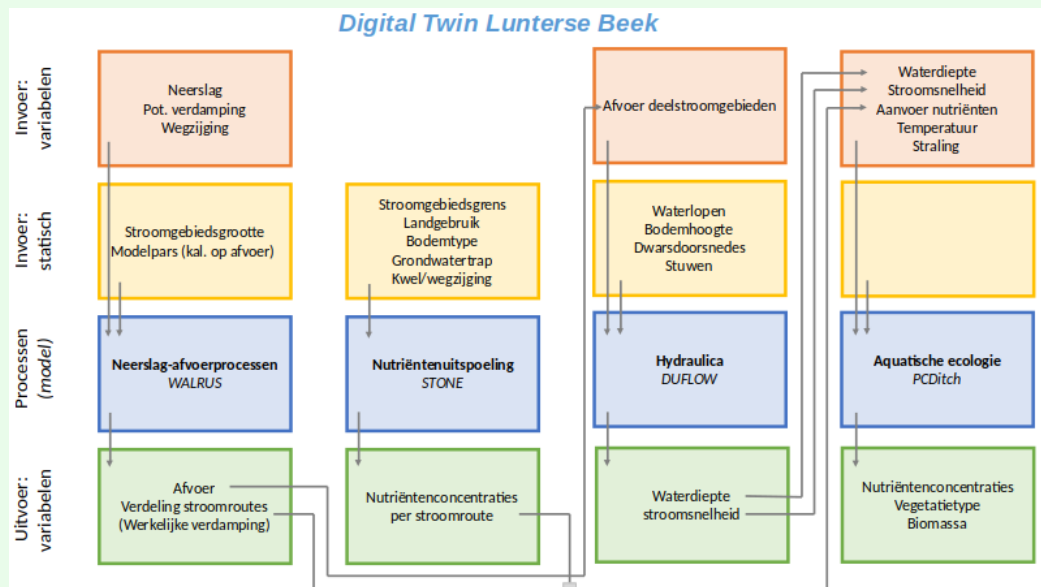
A conceptual schema should contain only those structural and behavioural aspects that are relevant to the universe of discourse. All aspects of physical external or internal data representation should be excluded. This requires the production of a conceptual schema, which is independent with respect to physical implementation technologies and platforms.

The goal of this essay is to make people developing digital twins of complex phenomena aware of the impact of their choices regarding resolution (in space, time and semantics) for the information that will be generated by the digital twin.

CASE

In the remainder of this essay, we will use as an example a Digital Twin set up by WUR for the 'Vallei en Veluwe' Water Board. This Digital Twin aims to integrate and process all information of the Lunterse Beek water system by modelling hydrology as well as water quality. The Digital Twin consists of a chain of four different software packages:

- WALRUS. The Wageningen Lowland RUNoff Simulator describes precipitation and runoff processes (as the name suggests, specifically for lowland areas) [Brauer et al. 2014a, b];
- STONE (Samen Te Ontwikkelen Nutriënten-Emissiemodel). Itself a chain of programs. The goal here is to model – at a national or regional scale – the fate of nutrients like nitrogen and phosphorus depending on policies and agricultural practices. This includes elements like manure excretion and distribution, NH_3 emission and deposition, N and P uptake by crops, transport and immobilization of N and P in soils, and leaching of N and P to surface and ground water [Wolf et al, 2003].
- DUFLOW. A program developed by STOWA and used in the Netherlands by water boards for one-dimensional hydraulic modelling of surface water. In the Digital Twin, it is used to calculate water quality aspects (such as algae growth) based on water movements and input loads for nutrients.
- PCDITCH. This component is modelling aquatic ecology of a water system [Janse, 2005¹², Van Gerwen 2016]. In this particular case the most compromising assumption is that the water body is supposed to be a rectangular body of non-moving water (a ditch), a quite inadequate description of the Lunterse Beek.



Digital Twins and data

A Digital Twin refers to the virtual replication of a physical ‘thing’ or ‘phenomenon’. This can be done either by using data directly, or by using models fed with data. The twin is dynamically updated by a continuous influx of new data. This sounds more simple than it is - data have multiple facets that are important if they are to be converted into information:

1. the data themselves
2. the metadata, *i.e.*, the information on the data, and
3. the uncertainty in the data, an aspect that is often not included.

Data (themselves)

Basically, the data form the primary material to work with. In general, the values are measured, observed or calculated and represent the value of a property.

Metadata

Metadata are data about data. Metadata should provide all you have to know to work with the data. They consists of ‘structured information that describes, explains, locates, or otherwise makes it easier to retrieve, use or manage an information resource, especially in a distributed network environment’ [Wikipedia]. Metadata normally are expressed by metadata elements and are usually categorized in a number of different types:

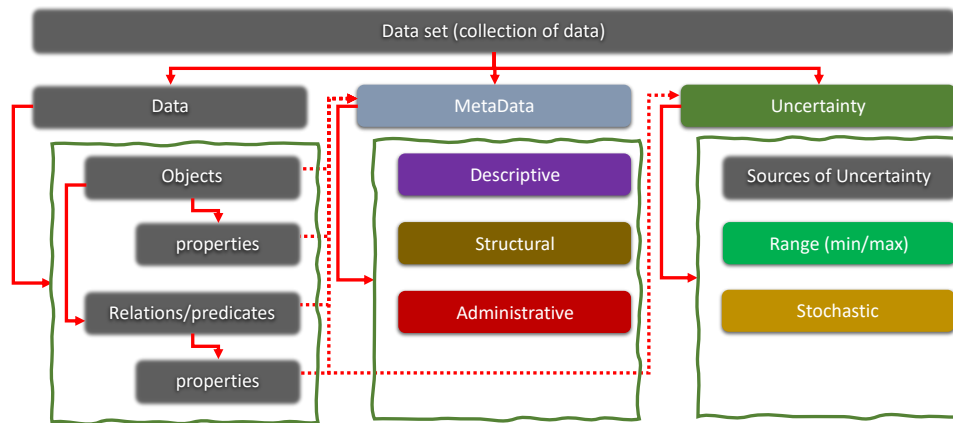
- Essential metadata: method of measurement, units of measurement, measurement/observation time, postprocessing methods, precision, location, etc.
- Descriptive metadata, describing an information resource for identification and retrieval through elements such as title, author, and abstract.
- Structural metadata, documenting relationships within and among objects through elements such as links to other components (*e.g.*, how pages are put together to form chapters). These include observation and measurement methods, units of measurement and other where relevant.
- Administrative metadata, for managing information resources through elements such as version number, archiving date, and other technical information for purposes of file management, rights management and preservation.

Metadata can be provided on the level of datasets, objects and properties, but all serve the same purpose. There are many standards the describe the metadata elements of that standard. To mention just two:

1. the ‘Dublin Core’ for networked resources and
2. for the spatial domain the ‘ISO 19115:2003 Geographic information -- Metadata standard’ defines how to describe geographical information and associated services, including contents, spatial-temporal purchases, data quality, access and rights to use.

Uncertainty

Data are usually the result of direct measurements, observations and derived as outcomes of transformations or calculations. Due to measurement errors, structural errors in our models, and errors in the estimated model parameters, a user of a Digital Twin will always be confronted with uncertainty. In Ten Broeke et al. (2021, in prep)¹³, more details are given on how to cope with uncertainty in a Digital Twin setting.



In short, the diagram above shows all aspects of data that are to be considered when working with data in any context. Only with information on all three components one can work with digital twins and cope with resolution in time, space and semantics that we will address in the next section.

Resolution in space, time and semantics of information generated by Digital Twins

Here we present an overview of resolution and accuracy in space, time and semantics of information generated by Digital Twins.

Space

In case spatial aspects play a role in a Digital Twin then spatial resolution should be considered for all components of that Digital Twin. The models that are part of the Digital Twin for the Lunterse Beek watershed¹⁴, were originally developed independently from each other. Each model is developed with its own specific (sometimes even implicit) spatial resolution and these will probably differ from each other. Spatial resolution concerns both input and output. Spatial resolution is often well defined for data, as is the case for remotely sensed imagery where each pixel has the same size (after post-processing by the data provider). For models like the STONE model, the spatial resolution is less clear. For STONE, the input data are given for areas called plots. A STONE-plot is a set of 250m x 250m raster cells not necessarily adjacent to each other, scattered over a region. Each STONE-plot resembles a unique combination of environmental variables (e.g., soil type, land use).

Other examples where spatial resolution is not necessarily constant over the modelling domain are finite element models (sometimes used for groundwater modelling). Finite element models usually have a higher resolution in regions where relatively steep gradients in model outputs are expected. But also, finite difference models may use a multiresolution discretisation when the boundary conditions of a finer grained local model are based on the output of an overlapping coarser grained regional model.

The resolution of model input is often identical to that of model output. However, this is not necessarily the case. Models may serve as aggregator by converting inputs at a finer resolution to outputs on a coarser resolution. An example is a simple runoff model where the inputs are defined for raster cells, and the output is a river outlet.

When there is a mismatch between the required resolution, and the available resolution, one may consider aggregation and disaggregation methods¹⁵. Aggregation methods convert data from a higher to a lower resolution whereas disaggregation methods do the reverse. Aggregation and disaggregation are sometimes also referred to as up- and downscaling, respectively. Obviously, data aggregation leads to loss of information, whereas data disaggregation leads to (locally) enrichment of information.

Data aggregation and disaggregation methods are different for data on continuous (real valued) variables and categorical variables (classes). For continuous variables, aggregated data are usually based on some weighted statistic, like weighted averages or medians. For categorical variables, the most dominant class is a useful aggregator. Aggregation and disaggregation of categorical variables clearly not only affect spatial resolution, but also semantic resolution (see below). Aggregation results in higher and disaggregation in lower levels of the semantic hierarchy.

The resolution at which Digital Twins internally operate, *i.e.*, the model resolution, is often more detailed than the resolution at which one wants to communicate results to users, *i.e.*, the presentation resolution. This means that the conversion of model resolution to presentation resolution also involves aggregation.

Time

Resolution in time or temporal resolution generally refers to the amount of time, the distance, it takes to retrace the equivalent spot in a cycle. In the case of measurements, a satellite may revisit the same place on Earth where a cycle is clearly present, for a sensor it is the frequency of its measurements where a cycle may be completely absent. For real world phenomena temporal resolution should relate to the dynamics of the phenomenon you are interested in. It is of no use measuring temperature once a month, if we are interested in temperature fluctuations within a day. It means that each phenomenon has its own rhythm they expose in the real world, which can depend on the hour of the day, the day of the year or the season in the year, etc. As indicated in the introduction, our Digital Twins represent complex systems, when looking closely this means that a complex system consists of a number of individual systems each with its own rhythm that react to each other. This only works if resolutions match in time by either adapting or fitting in the same pattern.

‘Pattern’ is an important concept in the definition for Rhythm. Patterns are present everywhere, rhythm is associated by humans mostly with time, but a pattern is a much wider concept, think of geometry and physical or biological phenomena that show a certain pattern or rhythm that can change again over time (transition to another rhythm).

The challenge to include rhythm when creating a Digital Twin as an abstraction of a real-world phenomenon next to identifying the right pattern is to also ‘catch’ the variation in that pattern. In most cases variations are seen as anomalies that cannot be constructed in computer programs generating patterns since they are not able to model imperfections. Apparently, phenomena in the real world are more erratic than a model can ever describe. A well-known quote from ‘weatherman’ Jan Pelleboer in the Netherlands is: ‘The weather forecast was good, but the weather did not adhere to it’, which indicates that despite the scientifically well-founded weather models including their interpretation it is not happening in reality. Compare it to music that is generated by a computer program that people perceive as too clean and less exciting. For Digital Twins, the most promising source to capture data is sensors that capture the real world, including the anomalies. However, of course it also still applies that noise and outliers must be recognized to prevent excessive deviations from reality.

In the past rhythm was evoked by people in their behaviour in our day-night cycle and determined by natural processes such as the seasons. Humans experienced this and responded within the limits of our observations. Nowadays we live in the digital era extending the limits of our observations. Data and information are easily accessible to support our decision making and anomalies can be used almost instantaneously as for example accidents causing traffic jams on our planned route to change to a more favourable route with less delays. Sharing rhythm, tuning rhythm, matching rhythm, and balancing rhythms as a measure for resolution are significant dynamics in Digital Twin/real-world processes and offer a new ground for decision making. Besides, rhythm as a concept that can improve replication into a Digital Twin, there are also general considerations to take into account. Scale and frequency are important ones that relate very much to the cycle or amplitude in a rhythm.

As for all data, data-quality should be an issue. Noise and outliers are always difficult to handle, especially dealing with rhythm but for sure they inform us on the imperfection and uncertainty of the Digital Twin. The concept of rhythm includes variation, so you could argue that noise is an essential part of rhythm presented by that pattern. And outliers? If not erroneous they can be messengers indicating changes in the rhythm. In that case we cannot ignore these data. For noise and outliers it should be assessed whether or not they are part of that rhythm or indeed errors.

The Digital Twin of the Lunterse Beek water system consists of four different software packages (computer models) each representing a different process in the system of the Lunterse Beek. Each of these processes has their own rhythm and consequently each of the software packages have a different temporal resolution that best suits the process it represents. One of the models (WALRUS) describes precipitation and runoff processes, these processes are determined by the weather and as we all know this changes a lot within a day. It uses inputs on an hourly basis, but smaller timesteps when necessary. The output data produced has a temporal resolution that represents a process that can affect river discharge uses a flexible timestep approach but can be within an hour. The STONE software models the fate of the nitrogen and phosphorous compounds. These processes are much slower and the data produced by these models have outputs that are aggregated by season. The DUFLOW and PCDITCH models calculate surface water quality aspects (such as algae growth) based on water movements and input loads of nutrients. These models have a temporal resolution of one day but use smaller timesteps if necessary. This Digital Twin is a system in which at least four processes are integrated each with different rhythms. The models that are combined in this Digital Twin have also different temporal resolutions which is fine as long as they are properly aligned.

Semantics

Categorical variables may be part of the Digital Twin as input variables, output variables or intermediate variables, not directly visible to the user. The notion of resolution, then, is related to the level of detail the variable contains. A simple toy example is a system recognizing cats and dogs in a set of pictures: simply saying 'cat' or 'dog' is less informative than also indicating what type of cat or dog is visible: e.g., 'Birman' or 'English shorthair' for types of cats, 'poodle' as an example of the category dog. What is interesting is that the categories form a hierarchy: cats as well as dogs are animals, mammals, have four legs, just to present a series of further categorizations. This is true for many other categorical variables as well. Obviously, if we can make predictions about the finest level of detail, then the Digital Twin is most useful: its output corresponds to the maximum amount of information. However, with more categories the risk of misclassifications quickly increases, and in most cases the number of training instances per category decreases. So clearly there is a trade-off, a level at which the Digital Twin presents maximally useful information at a certain level of reliability. Note that here we have been using the output of a system as an example; in principle the same considerations hold for input or intermediate variables. However, the concept of a hierarchy of classes is most often seen in output variables for the simple reason that most systems are not able to deal with hierarchies of factor levels. Simply presenting such a hierarchy to the user is less difficult: people intuitively know how to handle them. If the output of the procedure is directly channelled into another program, very often a translation step is needed to guarantee the correct hierarchical level of the new input. A final observation is that categories need not be mutually exclusive. The notion of probability of belonging to a class comes natural to many people, and results from a Digital Twin can be very informative when presented in such a manner. Fuzzy memberships are often converted to discrete – or 'crisp' - classifications in the final step of an analysis. In a hierarchy of classes this discretization basically occurs at every step when descending the ontology tree – the techniques mentioned below therefore only consider crisp classifications. In finding the balance between a detailed description and reliability, there are ways to go beyond the obvious: one can explicitly use the information in the hierarchy of categories to limit the Digital Twin outcomes to predictions of a certain trustworthiness, while achieving the highest level of detail, an approach known as hierarchical classification¹⁶. The key idea is that every class has a single, more general, parent class, and potentially many children. Three main approaches can be distinguished:

1. *Flat classifiers*: methods in this class will try to make predictions at the highest level of detail, and then propagate answers up the hierarchy. In the previous example, if the prediction would be 'poodle', then we know the parent node is 'dog'. In other words, information processing proceeds in a bottom-up fashion.
2. *Local classifiers*: these methods will create separate models for each node in the hierarchy of categories, and the hierarchy is traversed top-down.
3. Finally, *global classifiers* attempt to make predictions for the hierarchy as a whole.

Unsurprisingly, all three approaches have their difficulties. Flat classifiers need to be able to distinguish between a potentially very large number of classes. Again using the toy example to make the point: suppose we have twenty different types of dogs in our database and twenty types of cats, it is clear to see that a lot of effort will be spent trying to distinguish between the forty categories whereas only a few comparisons will be relevant. Information from non-leaf-level nodes (i.e., higher-level classes) is not considered at all. We already mentioned the fact that predictions at the highest level of detail are often hampered by a lack of data which can easily lead to overfitting.

Local classifiers avoid the issues of many classes and lack of detail-level data by starting at the most general level (where data are more abundant, and classes are few). So, in our example the first question would be 'Are you a man or a mouse' or rather 'cat or dog', which in many cases should not be too difficult to determine. Depending on the answer, the hierarchy will be descended along one of the branches. The main problem here is error propagation: an error at the top level will exclude a whole set of possible classes: often, the first steps carry too much weight and early decisions cannot be overturned easily. Of course, there are ways to mitigate these dangers: one could, e.g., descend only when there is enough confidence in the current decision, and only when there is enough data for the next decision to be made. However, this would lead to the need for extensive fine-tuning of the associated parameters, probably a road none of us would want to travel.

Unlike the previous two classes, global classifiers (sometimes also known as the 'big-bang approach') make predictions taking the whole hierarchy into account, and in that way potentially can avoid the aforementioned difficulties¹⁷¹⁸. One simple approach would be to use a distance-based error measure, where the distance is measured along the branches of the tree. In that way, a 'Birman' would be closer to an 'English short-hair' than to a 'Poodle'. To put it more bluntly: it is worse to mistake a horse for a skyscraper than for a donkey. Several strategies have been considered, e.g., based on clustering, trees or neural networks. In practice, however, global classifiers suffer from scalability issues, and are often computationally demanding.

Finding the right aggregation level

In a Digital Twin, where the goal of the analysis is not determined *a priori*, the model must be able to adapt to the question at hand and (des)aggregation becomes an essential part of presenting the results. This can be rather easy if the hierarchical structure is well-defined – in principle it could be just a matter of moving up or down in the hierarchy. In practice, however, many difficulties can be encountered. The user may not be able to decide what level is needed, for example, or the required level is not achievable with the data at hand.

It should be noted that the need to choose a good aggregation level is not unique to categorical variables - it is not uncommon for numerical variables, too. To name one example, quality labels are often given depending on the value of a numerical variable, where each quality class corresponds to a particular range of that variable. The 'resolution' of that numerical variable, or the precision with which that variable can be determined, then needs to be sufficient to unambiguously determine the correct class, but not higher

than that. So even here the required resolution of the analysis is determined by the question. In general, the secondary-school dogma 'only round numbers at the very last step of your calculation' seems to be appropriate here, too.

Of course, the aggregation level is not only important when presenting results – the process of generating the results in the first place crucially depends on what aggregation level the information is presented and can be processed. In the examples above, explicitly exploiting the hierarchical structure, aggregation can be automated to some extent. Local classifiers, for example, can be told to stop when there is not enough information to further separate the data into subcategories. Whether the result agrees with the requirements of the problem at hand remains to be seen – if the level of detail is larger than needed, aggregation can be performed. If it is not, only a partial answer is given.

Controlled vocabularies and resolution

With semantics, meaning is given to objects relevant within your domain of discourse. An ontology is a way to represent, formally name and define those object by categories, properties and relations between the concepts, data and entities that substantiate one, many, or all domains of discourse. More simply, an ontology is a way of showing the properties of a subject area and how they are related, by defining a set of concepts and categories that represent the subject.¹⁹

Ontologies are supported by controlled vocabularies and thesauri that provide a way to organize knowledge for subsequent retrieval. They are used in subject indexing schemes, subject headings, thesauri, taxonomies, code lists and other knowledge organization systems. Controlled vocabulary schemes mandate the use of predefined, authorised terms that have been preselected by the designers of the schemes, in contrast to natural language vocabularies, which have no such restriction²⁰. Controlled vocabularies, especially thesauri, contain, where applicable, parent-child relations that build up a hierarchy and provide means to select the proper resolution to be fit for purpose.

Vocabularies are organized and maintained per domain. As an example for the international Agro-domain AGROVOC is used as a vocabulary and for the Dutch 'water' domain the AQUO-lex²¹ is used in which in 2002 the previous used 'hydrologische woordenlijst' is fully included. Every term within the domain has a definition and where applicable a code, an acronym or abbreviation and dates indicating last updates and validity. Most important, it is maintained and users can propose requests for changes that will be evaluated according protocols agreed within the domain.

Implementation of those concepts based on 'meaning' facilitates the semantic web and with semantic technologies as the use of Resource Description Framework (RDF) and Web Ontology Language (OWL) machine readable processing of data is supported.

Decision-making with Digital Twins and dealing with uncertainty

To support efficient and efficacious decision-making, the presentation of the output of a Digital Twin is of paramount interest. Preferably the presentation and visualisation should support correct and effective processing of information by the decision-makers. This implies that for optimal presentation and visualisation insight is needed in the way decision-makers process information and utilise knowledge. An important factor is the accuracy of the output of a Digital Twin, which implies that the level of certainty is involved in decision-making. Therefore, not only the way of processing information is important, but also the way uncertainty is coped with.

The theoretical framework, used by Poortvliet et al.²² in an analysis of the utilisation of statistical information on uncertainty in flood risk management, might be a useful starting point to get insight in the way decision-makers process output from Digital Twins. This framework has two components: the way individuals process information, and the way uncertainty is coped with in cultures.

The first component of the theoretical framework concerns two different ways that can be distinguished in information processing by individuals: an analytical way and an experiential way^{23, 24, 25}. These two parallel modes can interact. The analytical way of information processing is based on reason, logic, equations, et cetera, also referred to as 'slow' since proper analysis takes time. The experiential way of information processing is based on affect, intuition, gut feeling et cetera, also referred to as 'fast'.

The second component of the theoretical framework concerns a determination of different strategies to cope with uncertainty, based on the 'monster metaphor' described by Smits^{26, 27} and applied in environmental sciences by Van der Sluijs²⁸. This metaphor is based on the idea that people use to order the world in mutually excluding categories, such as life and dead, human and animal, clean and unclean. Monsters can grow out of phenomena that belong to two categories that were considered as mutually excluding. The monster of uncertainty grows out from hybrids in the science-policy interface, such as Digital Twins. Coping strategies include *monster exorcism* (reducing uncertainty as much as possible), *monster embracement* (emphasizing uncertainty, for holistic or spiritual reasons or to perspective outcomes of scientific research), *monster adaptation* (e.g., use of statistically quantified uncertainty in a decision model), and *monster assimilation* (being transparent about uncertainty and accounting for multiple outcomes). In monster adaptation there is one optimal outcome that can be 'calculated', whereas in monster assimilation more outcomes are possible and a decision results from negotiations and balancing of interests. To these four coping strategies distinguished by Smits^{29, 30} and Van der Sluijs³¹ *monster denial* can be added (or at least neglection or downplaying³²).

Monster adaptation and monster assimilation appeal to the analytical or slow way of information processing, whereas the experiential or fast way of information processing can be related more to monster embracement or in denial, neglection or downplaying of the monster of uncertainty.

In a district water board, the various ways of processing information and coping with uncertainties can occur at various stages in the decision-making process. If efficiency and efficacy are aimed for, analytical processing of information and monster adaptation or monster assimilation are preferred. However, if opportunism rather than efficiency and efficacy prevails in decision-making, then experiential thinking and exorcism, denial, neglection or downplaying of the monster of uncertainty might predominate. This might be accompanied with persuasiveness, such as putting trust in research results by calling them plausible, and with a striving for consensus and support base among all parties in stake.

Assuming that efficiency and efficacy are aimed for in decision-making by the district water board 'Vallei en Veluwe', the aforementioned ways of information processing and strategies to cope with uncertainty are relevant for presenting Digital Twin results:

1. If high resolution goes with high uncertainty, interpretation of the Digital Twin results might be vulnerable to denial, neglect or downplaying of the uncertainty monster in an environment of experiential information processing. This hinders making efficient and efficacious decisions and thus should be discouraged.
2. If the resolution is higher than the scale of decision-making, e.g., for spatial resolution the size of a farm, catchment, water body, this might hinder analytical information processing since more complexity is involved than is needed to be fit for purpose.
3. If estimates, predictions and forecasts by Digital Twins are presented separately from information on their accuracy, this might invite to denial, neglect or downplaying of the uncertainty monster in an environment of experiential information processing. Therefore, integration of estimates, predictions or forecasts with information on uncertainty into one picture, is recommended.

In summary, we state that decisions to be supported rather than technical opportunities involved in a Digital Twin should determine the visualisation of Digital Twin results.

Communication with the Digital Twin

Probably one of the most important parts of a Digital Twin is communication and interaction with users. This is even more challenging when you realize that the questions you can ask the Digital Twin can be multiple. To understand the results, it means we want to see what is happening, to play with scenarios, to explore alternatives, to look at and compare results, etc. Communicating with a Digital Twin comes down to two main questions that are closely related: 'How to interact?' and 'How to visualise?'

How to interact?

A useful way to interact with a Digital Twin is via a dashboard. A dashboard is used to organise and present information of the Digital Twin and to provide controls for input and queries in an intuitive way. However, a dashboard for a Digital Twin goes beyond a graphical user interface (GUI) with pre-programmed buttons and sliders, as, although we know what the Digital Twin could provide, we don't know the exact questions end-user may ask, we need a far more flexible approach. The challenge here is to design a 'human' to 'machine' interface that supports both ends of the communication, sender to receiver and vice versa.

To ask questions to a 'machine' a possible solution is a query builder. A query builder can be used to compose a question for the Digital Twin in a simple, structured, and transparent way. The query builder translates the human readable question into instructions 'machine' coded to be processed by the Digital Twin.

Query builders are not new. Users of databases and geographical information systems are often familiar with query builders. Text-oriented query-builders generally use a domain specific language (DSL). Domain relates here to the technical implementation environment of the application, here the Digital Twin. This would be SQL for relational database systems, or JSON-based queries for document-oriented databases. A DSL can be as generic as those for databases, but also more specific like the DSL for the hydrological model Emerald³³. The challenge in the design here is to gap the bridge between DSL and the 'human' language of a user not directly familiar with the DSL. Examples of graphical query-builders are available in some geographical information systems, 3D-modelling software (www.blender.org), and flow-based programming software (e.g., nodered.org). A user builds a query by dragging nodes to a canvas and connecting these nodes by arrows (like building a mathematical graph with nodes and edges). Each node applies a specific action to the Digital Twin results in the dataflow. The dataflow is represented by the arrows connecting the nodes. These actions are for example denoted by verbs like select, filter, group, aggregate, arrange, plot, and tabulate.

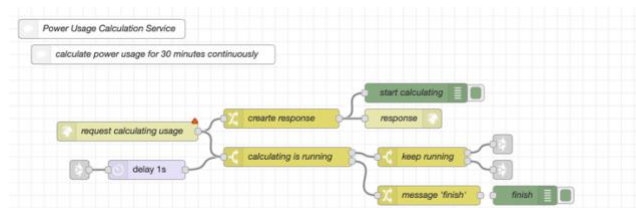


Figure 2 example nodered.org

Interactivity should also make 'what-if-scenario-analysis' possible by using interactive widgets to modify part of the original query. For example, for the Lunterse beek, this should answer questions like: what happens to nitrate loads in the Lunterse beek if dry periods get longer? What happens to baseflow? Does it drop below a critical threshold? How does the Lunterse beek respond to extreme precipitation events,

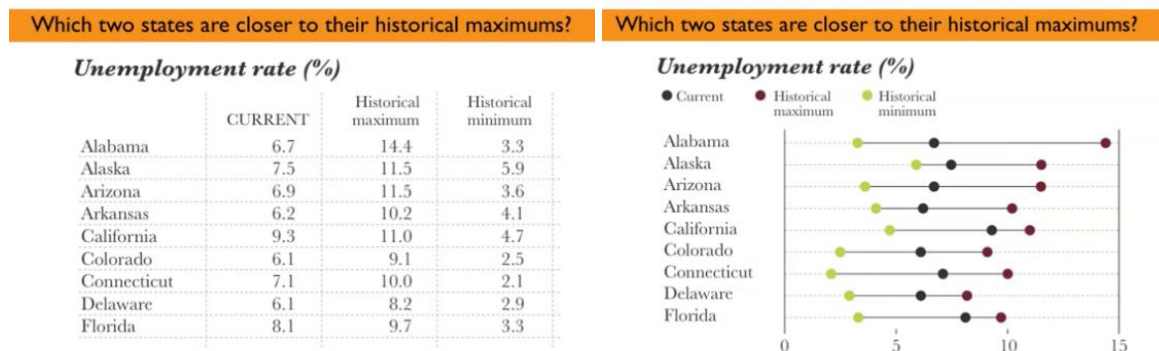
and how does that affect regions downstream? What measures are needed to retain more water in the watershed?

A dashboard also needs optimizers to answer questions in a reverse direction. E.g., what is the optimal weir-dimensions to guarantee vital summer discharge?

How to visualise?

After the query has been executed, the Digital Twin should return the requested information and present this information in easily digestible bite-size chunks. The information should preferably be presented in a graphical format. Graphics are very effective in communication, humans are visual able to properly interpret graphics so we need to take advantage of what a human brain can do. Cairo³⁴ stated that visual representations of data and phenomena should reveal things the bare eye would not be able to see otherwise. In visualizing information, we have to anticipate in the design what the user's brain will try to do. Data and information should be designed in a way the human brain can easily understand. 'Visualisation doesn't happen on a page or on a screen it happens in the viewers brain' as Robert Spence³⁵ stated. The way information is presented visually is targeted at the user and according to Kosslyn³⁶: 'the usefulness of a graph can be evaluated only in the context of the type of data, the questions the designer wants the readers to answer, and the nature of the audience'. The key to any visual design is the presentation of a cohesive, structured, readable and understandable image. Krug³⁷, an information architect and user experience professional on human-computer interaction and web usability phrased this as follows: 'Don't make me think'. A principle also used as a title for his book that became known as the 'usability bible' for designers, developers, webmaster and others.

Beside graphs there are of course other forms as maps, diagrams and tables. Tables, if really required, should be concise and interpretable at a glance. In the following example for the question asked, the left table is not suitable, but the visual on the right is.



Since interaction is important to communicate with the Digital Twin the graphics should be given as interactive dynamic visualisations. GUI-elements (widgets) like sliders, drop down lists, and gestures like panning, tapping/clicking, dragging, pinching, spreading and zooming should enhance and facilitate interactivity. For example, by clicking on a location in the map of the Lunterse beek, a time-series should show up with for instance discharge and nutrient load information on that location.

The main steps to design a good visualised project as for a Digital Twin are the following 1) Define focus, the story, the goals and the tasks; 2) do preliminary research on the data to visualised; 3) choose the graphic form that best suits the outcome of the first step; 4) make sketches and storyboards and structure the information; 5) complete the research and think about the visual style; 6) create the graphics, maps, and diagrams.

A Digital Twin is a virtual approximation of part of the real world. As such it introduces uncertainty. Ideally, when queried, the Digital Twin should not only provide state information of its real-world counterpart, but also quantify the associated uncertainty. This uncertainty information is important as it directly provides information on the approximation and provide means to interpret and use the results of the digital twin. If the phosphorus load in the Lunterse beek is afflicted with high uncertainty, users should be aware that they should use these results with caution. This information is also of importance for the Digital Twin developers as uncertainty information reveals potential caveats in their modelling system.

Another benefit of providing uncertainty information is that users can perform risk analysis. Risk is defined as the product of probability and impact. In particular, events that occur with high probability and have huge impacts need more attention.

Finally, as good practise uncertainty information should always be presented tightly connected to state information, to support the decision making also in the resolution best suited!

What to take into account when setting up a digital twin...

The goal of this essay is to make people who are developing digital twins of complex phenomena aware of the impact of their choices regarding resolution (in space, time and semantics) for the information that will be generated by the digital twin.

Whether it concerns the spatial, temporal or semantic domain, the resolution at which the analysis is performed and that is used for presenting the results to the user should be carefully chosen – this is perhaps the central point of this essay. In the temporal domain, the resolution of the model depends not only on the resolution of the data, but also on the time scale of the events in the real-world, or in case of cyclic phenomena, the rhythm of the real-world twin half. Within a complex system several processes with often different rhythms and resolutions are interacting which should be reflected in the Digital Twin. Similarly, if spatial components are present their resolution should match the data as well as the spatial variability of the modelled processes. Also here it is the rule rather than the exception that different resolutions play a part at different stages. In the semantic dimension same applies: one should think about the level of detail at which analyses are performed and data are used. However, the most important decision perhaps in all three cases is the resolution at which results are presented to the user, and the level of uncertainty or confidence associated with that. There, the resolution should match the question – and since this in a Digital Twin is typically undefined during the implementation, the Digital Twin should be flexible, i.e., have the possibility to adapt the chosen answer resolution to the situation at hand. It also follows that a decision about this resolution needs to be taken within the Digital Twin itself, without human intervention – this is not an easy task which already comes close to Artificial Intelligence (AI). An alternative is to provide the user with UI tools to adjust the resolution of the answer to be optimally useful. In that case it is of utmost importance, as stressed in the preceding paragraphs, to also include estimates of uncertainty, to avoid false confidence in the results. Last but not least it is important to question the efficiency of the Digital Twin during the implementation process: when do you stop investing in improving accuracy and resolution and start optimising and quantifying uncertainty. It remains a trade-off between value of perfect information and expected value of including uncertainty.

The following set of statements reflects the key messages that are discussed in the essay.

- To ensure useful answers to questions that are not yet known a careful set up of the Digital Twin is required. One of the most important factors is the level of detail or resolution that is required for the abstraction of the real-world case.
- Whether it concerns the spatial, temporal or semantic domain, the resolution at which the analysis is performed and the resolution that is used for presenting the results to the user should be carefully chosen.
- In the temporal domain, the resolution of the model depends not only on the resolution of the data, but also on the time scale of the events in the real-world, or in case of cyclic phenomena, the rhythm of the real-world twin half.
- Within a complex system several processes with often different rhythms and resolutions are interacting which should be reflected in the Digital Twin
- if spatial components are present their resolution should match the data as well as the spatial variability of the modelled processes within the Digital Twin
- The resolution used for presenting the results of a Digital Twin should match the question posed to the Digital Twin. Since it is typically undefined during the implementation what questions will be asked to a Digital Twin, the Digital Twin should be flexible, i.e., have the possibility to adapt the chosen answer resolution to the situation at hand.

- In general, a higher resolution is a higher level of detail, but not necessarily better or worse, depending on the intended use and the associated uncertainty.
- When there is a mismatch between the required resolution, and the available resolution, one may consider aggregation and disaggregation methods.
- A Digital Twin is a virtual approximation of part of the real world. As such it introduces uncertainty. Ideally, when queried, the Digital Twin should not only provide state information of its real world counterpart, but also quantify the associated uncertainty. This uncertainty information is important as it directly provides information on the approximation and provide means to interpret and use the results of the digital twin.
- Take into account the trade-off between value of perfect information and expected value of including uncertainty during implementation

References

- ¹ Grieves M (2014) Digital twin: Manufacturing excellence through virtual factory replication. White paper, Florida Institute of Technology.
- ² Tomko and Winter, 2018. Beyond digital twins –A commentary. In Environment and Planning B: Urban Analytics and City Science 2019, Vol. 46(2) 395–399
- ³ Roberto Poli, 2013: in ‘Cadmus, leadership in thought that leads to action’, Volume 2 - Issue 1, October 2013
- ⁴ Walsh J. 8 Myths About Digital Twins Exposed—Here’s the Reality 2020. https://www.engineering.com/DesignSoftware/DesignSoftwareArticles/ArticleID/20561/8-Myths-About-Digital-Twins-ExposedHeres-the-Reality.aspx?utm_source=facebook&utm_medium=social&utm_campaign=facebook (accessed September 11, 2020).
- Walsh J, Tolle D. ASSESS-Theme-Positioning-Paper-Digital-Twins-1.10.pdf. ASSESS Initiative; 2020.
- ⁵ Savian C. Digital twins for a safer Built Environment 2019. <https://www.theiet.org/impact-society/sectors/built-environment/built-environment-blog-posts/digital-twins-for-a-safer-built-environment/> (accessed October 12, 2020)
- ⁶ [Taken from position paper DT-Methodology Platform, WUR 2020 in prep, where the rules are described in this paper more in detail].
- ⁷ The Open Data Institute. R&D: Digital twins. RD Digit Twins n.d. <https://theodi.org/project/rd-digital-twins/> (accessed October 9, 2020).
- ⁸ Tomasz Kielar, 2019, Digital Twin Explained – The Next Thing After IoT, blog of May 8, 2019 on <https://www.pgs-soft.com/blog/digital-twin-explained-the-next-thing-after-iot/>
- ⁹ Shunlin Liang, Xiaowen Li, Jindi Wang (editors), Chapter 1 - A Systematic View of Remote Sensing, Editor(s):, Advanced Remote Sensing, Academic Press, 2012, Pages 1-31, ISBN 9780123859549, <https://doi.org/10.1016/B978-0-12-385954-9.00001-0>. (<http://www.sciencedirect.com/science/article/pii/B9780123859549000010>)
- ¹⁰ Ewa Dębińska and Piotr Cichociński 2007; Conceptual Modelling of Real Estates for the Purposes of Mass Appraisal; in fig proceedings 2007 (https://www.fig.net/resources/proceedings/fig_proceedings/fig2007/papers/ts_4d/ts04d_01_debinska_cichocinski_1290.pdf)
- ¹¹ INSPIRE data specifications D2.5 Generic conceptual model; 2013; Drafting team data specifications (https://inspire.ec.europa.eu/documents/Data_Specifications/D2.5_v3.4rc3_vs_3.4rc2.pdf)
- ¹² Janse, J. H. (2005). Model studies on the eutrophication of shallow lakes and ditches. Dissertatie, Universiteit Wageningen, <http://edepot.wur.nl/121663> (PDF)
- ¹³ Ten Broeke, G., H. Stigter & R. Wehrens, 2021 (in prep). Onzekerheid in Digital Twins: inventarisatie, doorwerking en beperking.

-
- ¹⁴ Ten Broeke, G., H. Stigter & R. Wehrens, 2021 (in prep). Onzekerheid in Digital Twins: inventarisatie, doorwerking en beperking.
- ¹⁵ Walvoort, D. J. J., Knotters, M. & van Egmond, F. M., 2019, Interpolatie, aggregatie en desaggregatie van ruimtelijke bodemgegevens in de Basisregistratie Ondergrond (BRO). Wageningen: Wettelijke Onderzoekstaken Natuur & Milieu. 48 p. (WOt-technical report; no. 167)
- ¹⁶ Data Mining and Knowledge Discovery, 22:31–72.
- ¹⁷ Kosmopoulos, A., Partalas, I., Gaussier, E., Paliouras, G., & Androutsopoulos, 2015. "Evaluation measures for hierarchical classification: a unified view and novel approaches", Data Mining and Knowledge Discovery, 29(3), 820-865
- ¹⁸ Sokolova, M., Lapalme, G., 2009, "A systematic analysis of performance measures for classification tasks", Information Processing and Management 45(4), 427–437.
- ¹⁹ [https://en.wikipedia.org/wiki/Ontology_\(information_science\)](https://en.wikipedia.org/wiki/Ontology_(information_science))
- ²⁰ https://en.wikipedia.org/wiki/Controlled_vocabulary
- ²¹ <https://www.aquo.nl/index.php/Categorie:Begrippen>
- ²² Poortvliet, P.M., M. Knotters, P. Bergsma, J. Verstoep and J. van Wijk, 2019. On the communication of statistical information about uncertainty in flood risk management. Safety Science 118: 194-204. <https://doi.org/10.1016/j.ssci.2019.05.024>
- ²³ Epstein, S., 1994. Integration of the cognitive and the psychodynamic unconscious. American Psychologist 49(8): 709-724.
- ²⁴ Slovic, P., M.L. Finucane, E. Peters and D.G. MacGregor, 2004. Risk as analysis and risk as feelings: some thoughts about affect, reason, risk, and rationality. Risk Analysis 4(2): 311-322.
- ²⁵ Kahneman, D., 2011. Thinking, fast and slow. New York, Farrar, Straus and Giroux.
- ²⁶ Smits, M., 2002. Monsterbezwinging: De culturele domesticatie van nieuwe technologie. Proefschrift, Universiteit Twente.
- ²⁷ Smits, M., 2006. Taming monsters: the cultural domestication of new technology. *Technology in Society* 28: 489-504.
- ²⁸ Van der Sluijs, J., 2005. Uncertainty as a monster in the science-policy interface: four coping strategies. *Water Science & Technology* 52(6): 87-92.
- ²⁹ Smits, M., 2002. Monsterbezwinging: De culturele domesticatie van nieuwe technologie. Proefschrift, Universiteit Twente.
- ³⁰ Smits, M., 2006. Taming monsters: the cultural domestication of new technology. *Technology in Society* 28: 489-504.
- ³¹ Van der Sluijs, J., 2005. Uncertainty as a monster in the science-policy interface: four coping strategies. *Water Science & Technology* 52(6): 87-92.

³² Poortvliet, P.M., M. Knotters, P. Bergsma, J. Verstoep and J. van Wijk, 2019. On the communication of statistical information about uncertainty in flood risk management. *Safety Science* 118: 194-204.
<https://doi.org/10.1016/j.ssci.2019.05.024>

³³ Walvoort, D.J.J & Bierkens, M.F.P., 1999. Emerald: A Stochastic Modelling Approach for Rapid Assessment of Groundwater Dynamics. DLO Winand Staring Centre report 171.

³⁴ Cairo, Alberto. *The Functional Art : an Introduction to Information Graphics and Visualization*. Berkeley, CA :New Riders, 2013.

³⁵ Robert Spence. 2007. *Information Visualization: Design for Interaction* (2nd Edition). Prentice-Hall, Inc., USA.

³⁶ Kosslyn, S. M. (2006). *Graph design for the eye and mind*. Oxford University Press.
<https://doi.org/10.1093/acprof:oso/9780195311846.001.0001>

³⁷ Krug, Steve. *Don't Make Me Think, Revisited: a Common Sense Approach to Web Usability*. [Berkeley, Calif.] :New Riders, 2014.